

GATE Data Science and Artificial Intelligence

Sample Paper – 8

Duration: 128 Minutes

Maximum Marks: 72

Instructions

- This paper contains **46 questions** covering the core **Data Science and Artificial Intelligence** subject of the GATE paper conducted by the IITs and IISc (General Aptitude and Engineering Mathematics are provided as separate papers).
- **Q1–Q20 carry 1 mark each and Q21–Q46 carry 2 marks each**, for a total of **72 marks**.
- Each question carries **negative marking of one-third of its allotted marks**: a wrong 1-mark question costs **0.33** and a wrong 2-mark question costs **0.67**. Unattempted questions carry no penalty.
- Every question here is a single-correct **MCQ**. The real GATE paper also has Multiple Select (MSQ) and Numerical Answer Type (NAT) questions, and **those carry no negative marking**.
- Attempt this paper in one timed sitting of about **128 minutes**, the share this core subject earns at GATE's overall pace of 180 minutes for 65 questions. All logarithms in information-theoretic quantities (entropy, information gain) are to base 2 and expressed in bits unless stated otherwise.

Probability & Statistics

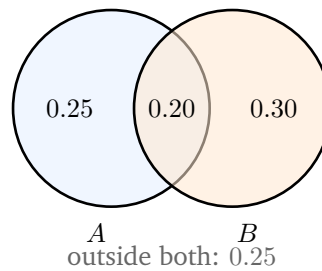
- Q1.** A bag contains 3 red and 5 blue balls. Two balls are drawn one after the other *without* replacement. The probability that *both* drawn balls are blue is: [1 mark]

- (A) $\frac{25}{64}$
(B) $\frac{3}{28}$
(C) $\frac{5}{14}$

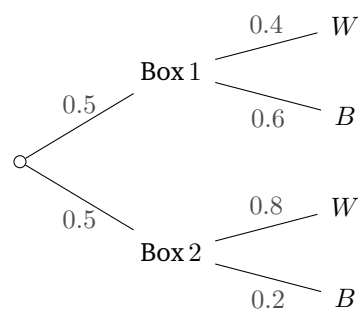


(D) $\frac{5}{8}$

- Q2.** For two events A and B the probabilities of the disjoint regions are shown on the Venn diagram below (the region outside both circles has probability 0.25). The probability $P(A \cup B)$ is: [1 mark]



- (A) 0.75
 (B) 0.45
 (C) 0.20
 (D) 0.55
- Q3.** Two boxes are given: Box 1 holds 2 white and 3 black balls, and Box 2 holds 4 white and 1 black ball. A box is chosen at random (each with probability 0.5) and one ball is drawn, as shown in the probability tree. Given that the drawn ball is *white*, the probability that it came from Box 1 is: [1 mark]

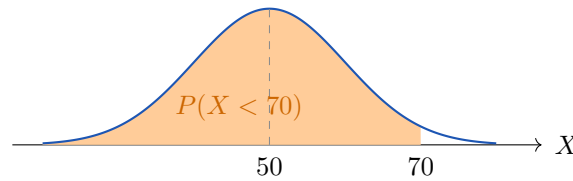


- (A) 0.50
 (B) $\frac{1}{3}$
 (C) $\frac{2}{3}$



(D) 0.40

- Q4.** A feature X is normally distributed with mean $\mu = 50$ and standard deviation $\sigma = 10$. Using the shaded left region of the standard normal curve shown, the probability $P(X < 70)$ is nearest to: [1 mark]



- (A) 0.0228
(B) 0.5000
(C) 0.9772
(D) 0.8413
- Q5.** Defects occur along a production line following a Poisson process at an average rate of 2 defects per metre of cloth. The probability of finding *exactly zero* defects in a given one-metre length is nearest to (take $e^{-2} = 0.1353$): [1 mark]
- (A) 0.2707
(B) 0.8647
(C) 0.4060
(D) 0.1353
- Q6.** A discrete random variable X takes the values $-1, 0, 1$ with probabilities 0.3, 0.4, 0.3 respectively. The value of $E[X^2]$ is: [1 mark]

- (A) 0.0
(B) 0.3
(C) 0.6
(D) 1.0



Q7. A sample of size $n = 64$ drawn from a population with known standard deviation $\sigma = 16$ has sample mean $\bar{x} = 40$. For a 95% confidence interval for the population mean (using $z_{0.025} = 1.96$), the margin of error (half-width of the interval) is:

[1

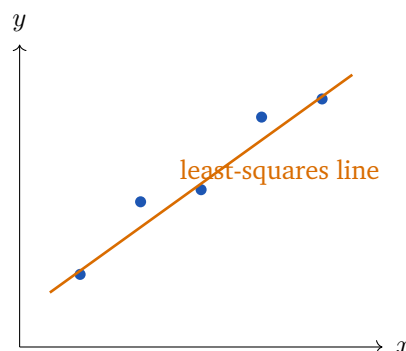
mark]

- (A) 1.96
- (B) 16.00
- (C) 2.00
- (D) 3.92

Q8. For a bivariate data set the sums of squares and cross-products about the means are $S_{xx} = 25$, $S_{yy} = 100$ and $S_{xy} = 30$, and the scatter with its least-squares line is shown. The Pearson correlation coefficient r between x and y is:

[1

mark]



- (A) 1.20
- (B) 0.30
- (C) 0.36
- (D) 0.60

Q9. A fair coin is tossed 5 times independently. The probability of obtaining exactly 2 heads is:

[1 mark]

- (A) 0.5000
- (B) 0.15625



(C) 0.3125

(D) 0.2500

Q10. The time until the next request arrives at a web server follows an exponential distribution with a mean of 4 seconds. The probability that the next request arrives *within* 4 seconds is (take $e^{-1} = 0.3679$): [1 mark]

(A) 0.368

(B) 0.632

(C) 0.135

(D) 0.865

Linear Algebra, Calculus & Optimization

Q11. Consider the matrix $A = \begin{bmatrix} 1 & 2 & 1 \\ 2 & 4 & 2 \\ 3 & 6 & 3 \end{bmatrix}$. The rank of A is: [1 mark]

(A) 1

(B) 2

(C) 3

(D) 0

Q12. The determinant of the matrix $B = \begin{bmatrix} 1 & 2 & 0 \\ 0 & 3 & 1 \\ 2 & 0 & 1 \end{bmatrix}$ is: [1 mark]

(A) 3

(B) 7

(C) 11

(D) 5

Q13. The *smaller* eigenvalue of the matrix $C = \begin{bmatrix} 2 & 2 \\ 1 & 3 \end{bmatrix}$ is: [1 mark]

(A) 4



- (B) 1
- (C) 5
- (D) 2

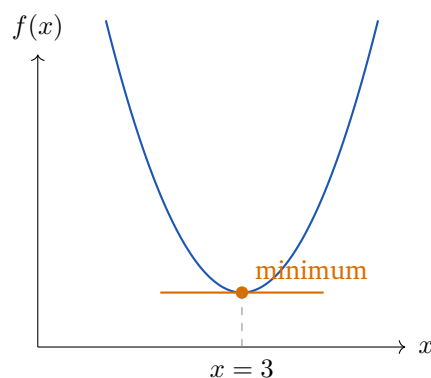
Q14. For a real matrix A , the matrix $A^T A$ has eigenvalues 16 and 9. The Frobenius norm $\|A\|_F = \sqrt{\text{tr}(A^T A)}$ is: [1 mark]

- (A) 5
- (B) 25
- (C) 7
- (D) 4

Q15. For the vector $v = (6, 8)$, the Euclidean (ℓ_2) norm $\|v\|_2$ is: [1 mark]

- (A) 14
- (B) 10
- (C) 100
- (D) 48

Q16. The function $f(x) = x^2 - 6x + 13$ is minimized over the real line, as suggested by its convex graph below. The minimum value of f is: [1 mark]

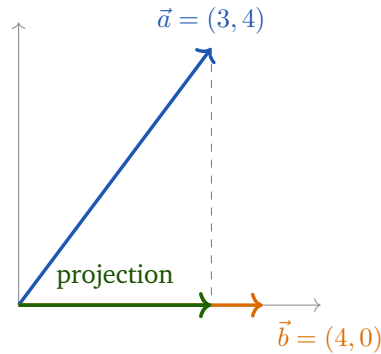


- (A) 3
- (B) 13
- (C) 4



(D) 0

Q17. The scalar projection (component) of $\vec{a} = (3, 4)$ onto $\vec{b} = (4, 0)$, as shown, is: [1 mark]



(A) 4

(B) 3

(C) 5

(D) 7

Q18. Gradient descent is used to minimize $f(x) = (x-3)^2$ with a fixed learning rate $\eta = 0.2$. Starting from $x_0 = 0$, the value of x_1 after one update $x_1 = x_0 - \eta f'(x_0)$ is: [1 mark]

(A) 1.2

(B) -1.2

(C) 0.6

(D) 6.0

Q19. Consider the function $f(x, y) = 2x^2 + 3y^2 + xy$ on $\mathbb{R}^2$. Based on its Hessian, the function is: [1 mark]

(A) concave everywhere

(B) neither convex nor concave

(C) convex everywhere

(D) linear (affine)



Programming, Data Structures & Algorithms

Q20. What is the output of the following Python snippet?

```
a = [1, 2, 3, 4, 5]
print(a[::-2])
```

[1 mark]

- (A) [2, 4]
- (B) [1, 3, 5]
- (C) [1, 2, 3]
- (D) [5, 3, 1]

Q21. What is the output of the following Python snippet?

```
x = "hello"
print(x[::-1])
```

[2 marks]

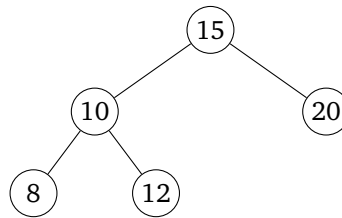
- (A) olleh
- (B) hello
- (C) hell
- (D) raises an error

Q22. The following operations are performed on an initially empty *queue*, in order: enqueue(1), enqueue(2), dequeue(), enqueue(3), enqueue(4), dequeue(), dequeue(). The value returned by the *last* dequeue() is: [2 marks]

- (A) 1
- (B) 4
- (C) 3
- (D) 2

Q23. Starting from an empty binary search tree, the keys 15, 10, 20, 8, 12 are inserted in the given order, producing the tree shown. The number of *leaf* nodes in the resulting tree is: [2 marks]





- (A) 2
- (B) 3
- (C) 4
- (D) 5

Q24. A hash table of size 7 uses the hash function $h(k) = k \bmod 7$ with *linear probing* for collision resolution. The keys 15, 22, 8 are inserted in this order into an otherwise empty table. The final index occupied by key 8 is: [2 marks]

- (A) 1
- (B) 2
- (C) 3
- (D) 8

Q25. Binary search is performed on a sorted array of 63 distinct elements. The maximum number of key comparisons required in the worst case is: [2 marks]

- (A) 63
- (B) 32
- (C) 6
- (D) 7

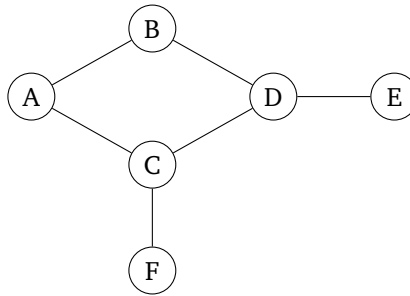
Q26. The worst-case time complexity of the standard Merge Sort algorithm on an array of n elements is: [2 marks]

- (A) $O(n^2)$
- (B) $O(n)$



- (C) $O(\log n)$
- (D) $O(n \log n)$

Q27. A depth-first search (DFS) is started from vertex A in the undirected graph shown, always visiting adjacent vertices in increasing alphabetical order. The order in which the vertices are first visited is: [2 marks]



- (A) A, B, C, D, E, F
- (B) A, C, B, D, F, E
- (C) A, B, D, E, C, F
- (D) A, B, D, C, F, E

Q28. The running time of an algorithm satisfies the recurrence $T(n) = 4T(n/2) + n$, with $T(1) = 1$. By the Master theorem, $T(n)$ is: [2 marks]

- (A) $\Theta(n^2)$
- (B) $\Theta(n \log n)$
- (C) $\Theta(n)$
- (D) $\Theta(n^2 \log n)$

Database Management & Warehousing

Q29. Union-compatible relations R and S contain 8 tuples and 5 tuples respectively. The *maximum* possible number of tuples in $R \cup S$ is: [2 marks]

- (A) 13
- (B) 40
- (C) 8



(D) 3

Q30. A table T has 6 rows, and its column c holds the values 10, 10, 20, NULL, 30, 20 (one value per row). What does the query `SELECT COUNT(DISTINCT c) FROM T;` return? [2 marks]

(A) 6

(B) 5

(C) 3

(D) 4

Q31. A relation $R(A, B, C)$ has the functional dependencies $AB \rightarrow C$ and $C \rightarrow B$. Its candidate keys are AB and AC . The highest normal form that R satisfies is: [2 marks]

(A) BCNF

(B) 3NF

(C) 2NF

(D) 1NF

Q32. Relation $R(A, B)$ contains the tuples $\{(1, p), (2, p), (3, q)\}$ and relation $S(B, C)$ contains $\{(p, x), (p, y), (q, z)\}$. The number of tuples in the natural join $R \bowtie S$ (on the common attribute B) is: [2 marks]

(A) 3

(B) 6

(C) 9

(D) 5

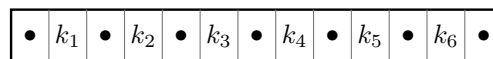
Q33. In OLAP, the operation that selects a single value for one dimension of a data cube (for example, fixing $\text{Year} = 2023$), thereby reducing the cube's dimensionality by one, is called: [2 marks]

(A) roll-up



- (B) drill-down
- (C) dice
- (D) slice

Q34. The figure shows the structure of an internal node of a B^+ tree of order $m = 7$ (each internal node holds up to m child pointers, drawn as $\bullet$, separated by keys). The maximum number of keys an internal node of this tree can hold is: [2 marks]



up to 7 pointers ($\bullet$) and 6 keys

- (A) 6
- (B) 7
- (C) 5
- (D) 3

Q35. A relation has attributes $\{A, B, C, D, E\}$ with functional dependencies $A \rightarrow B$ and $B \rightarrow C$. The attribute closure A^+ is: [2 marks]

- (A) $\{A, B, C\}$
- (B) $\{A, B, C, D, E\}$
- (C) $\{A\}$
- (D) $\{A, B\}$

Machine Learning & Artificial Intelligence

Q36. A logistic regression model has weight vector $w = (2, -1)$ and bias $b = 0$, and uses the sigmoid $\sigma(z) = 1/(1 + e^{-z})$. For the input $x = (1, 2)$, the predicted probability $\sigma(w^T x + b)$ is: [2 marks]

- (A) 0.73
- (B) 0.88
- (C) 0.27



(D) 0.50

Q37. A very simple model attains *low* accuracy on the training set *and* similarly low accuracy on unseen test data. In terms of the bias–variance decomposition, this behaviour (underfitting) is best described as: [2 marks]

- (A) high bias, low variance
- (B) low bias, high variance
- (C) low bias, low variance
- (D) high bias, high variance

Q38. A 1-nearest-neighbour classifier (Euclidean distance) is trained on the labelled points $A(1, 1)$: Class 0, $B(5, 5)$: Class 1, $C(2, 2)$: Class 0, $D(6, 6)$: Class 1. The query point $Q(4, 4)$ is classified as: [2 marks]

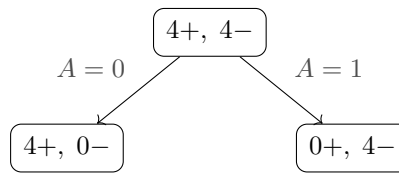
- (A) Class 1
- (B) Class 0
- (C) a tie (undecided)
- (D) cannot be determined

Q39. A training set for a binary classification task contains 10 examples, of which 5 are positive and 5 are negative. The entropy of this set (in bits) is: [2 marks]

- (A) 0.5
- (B) 1.0
- (C) 0.811
- (D) 0.0

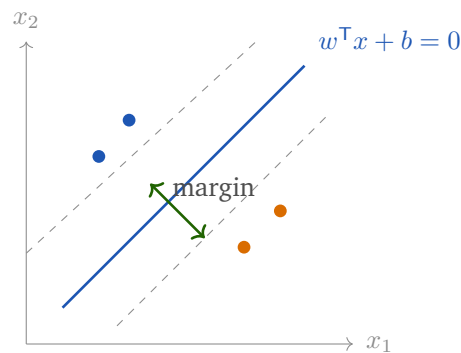
Q40. A node in a decision tree contains 4 positive and 4 negative examples. A candidate attribute A splits it into two children as shown: $A = 0$ gives $(4+, 0-)$ and $A = 1$ gives $(0+, 4-)$. The information gain of this split is: [2 marks]





- (A) 1.0
- (B) 0.5
- (C) 0.189
- (D) 0.0

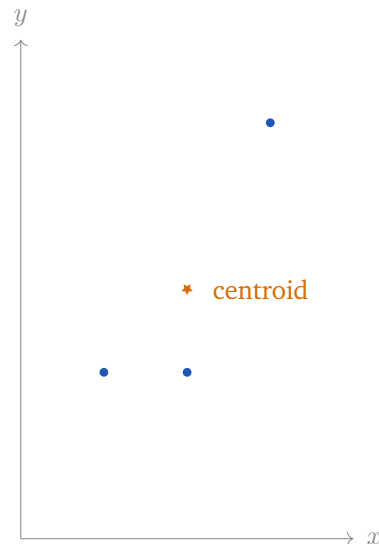
Q41. For a linear support vector machine the separating hyperplane $w^T x + b = 0$ is shown with its two margin boundaries. If $\|w\| = 2$, the width of the margin, equal to $2/\|w\|$, is: [2 marks]



- (A) 4
- (B) 2
- (C) 0.5
- (D) 1

Q42. In one iteration of the k -means algorithm, a cluster is assigned the three points $(2, 4)$, $(4, 4)$ and $(6, 10)$, shown below. The updated centroid of this cluster is: [2 marks]



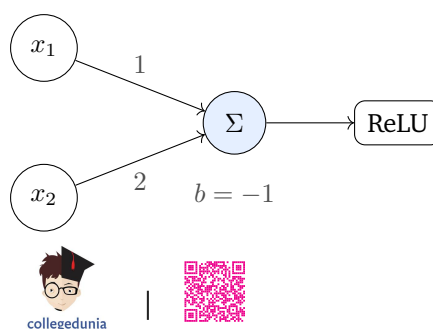


- (A) (4, 4)
- (B) (4, 6)
- (C) (6, 4)
- (D) (12, 18)

Q43. Principal component analysis is applied to a data set whose covariance matrix has eigenvalues 6, 2, 1, 1. The fraction of the total variance explained by the *first* principal component alone is: [2 marks]

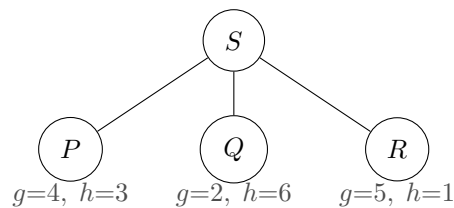
- (A) 80%
- (B) 20%
- (C) 40%
- (D) 60%

Q44. A single neuron with a ReLU activation computes $\text{ReLU}(w_1x_1 + w_2x_2 + b)$, as shown, where $\text{ReLU}(z) = \max(0, z)$. With inputs $x_1 = 2$, $x_2 = 1$, weights $w_1 = 1$, $w_2 = 2$ and bias $b = -1$, the output of the neuron is: [2 marks]



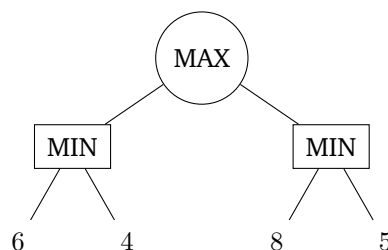
- (A) 0
- (B) 3
- (C) 5
- (D) 1

Q45. A* search maintains three frontier nodes P , Q and R with g (path cost so far) and h (heuristic) values shown in the tree below, where $f = g + h$. The node A* selects for expansion next is: [2 marks]



- (A) P
- (B) Q
- (C) a tie (all equal)
- (D) R

Q46. In the two-ply game tree shown, the root is a MAX node and the second level are MIN nodes; the leaf values are given. The minimax value backed up to the root is: [2 marks]



- (A) 4
- (B) 6
- (C) 5
- (D) 8



Detailed Solutions**Q1.****Solution**

Concept – probability without replacement: Multiply the probability of the first blue by the conditional probability of the second blue.

Step 1 – count the balls: Total = $3 + 5 = 8$ balls, of which 5 are blue.

Step 2 – probability the first ball is blue: $P(\text{1st blue}) = \frac{5}{8}$

Step 3 – probability the second is blue given the first was blue: Now 4 blue remain out of 7 balls: $P(\text{2nd blue} \mid \text{1st blue}) = \frac{4}{7}$

Step 4 – multiply: $P(\text{both blue}) = \frac{5}{8} \times \frac{4}{7} = \frac{20}{56}$

Step 5 – simplify: $\frac{20}{56} = \frac{5}{14}$

Final Answer: $\frac{5}{14} \Rightarrow \boxed{\text{C}}$

Answer: (C) [Go Back to Q1](#)

Q2.**Solution**

Concept – union from disjoint Venn regions: $P(A \cup B)$ is the sum of the “only A”, “only B” and “both” regions.

Step 1 – read the three regions inside the circles: Only A = 0.25, both = 0.20, only B = 0.30.

Step 2 – add them: $P(A \cup B) = 0.25 + 0.20 + 0.30$

Step 3 – evaluate: $P(A \cup B) = 0.75$

Step 4 – sanity check the whole sample space: $0.75 + 0.25$ (outside) = 1.00, so the diagram is consistent.

Final Answer: $P(A \cup B) = 0.75 \Rightarrow \boxed{\text{A}}$

Answer: (A) [Go Back to Q2](#)



Q3.

Solution

Concept – Bayes' theorem (reverse the tree): $P(B_1 | W) = \frac{P(B_1) P(W | B_1)}{P(B_1) P(W | B_1) + P(B_2) P(W | B_2)}$.

Step 1 – white-ball probability from each box: Box 1: $\frac{2}{5} = 0.4$; Box 2: $\frac{4}{5} = 0.8$.

Step 2 – probability from the Box 1 branch: $P(B_1) P(W | B_1) = 0.5 \times 0.4 = 0.20$

Step 3 – probability from the Box 2 branch: $P(B_2) P(W | B_2) = 0.5 \times 0.8 = 0.40$

Step 4 – total probability of drawing white: $P(W) = 0.20 + 0.40 = 0.60$

Step 5 – apply Bayes' theorem: $P(B_1 | W) = \frac{0.20}{0.60} = \frac{1}{3}$

Final Answer: $P(B_1 | W) = \frac{1}{3} \Rightarrow \boxed{\text{B}}$

Answer: (B) [Go Back to Q3](#)

Q4.

Solution

Concept – standardize, then use the normal table: Convert X to the standard normal $Z = \frac{X - \mu}{\sigma}$.

Step 1 – standardize the cut-off $X = 70$: $Z = \frac{70 - 50}{10}$

$$Z = \frac{20}{10} = 2$$

Step 2 – write the required probability in Z : $P(X < 70) = P(Z < 2)$

Step 3 – use the standard normal value: $P(Z < 2) = 0.9772$

Step 4 – interpret the shaded region: The shaded area is the whole left region up to $Z = 2$, which is 0.9772.

Final Answer: $P(X < 70) = 0.9772 \Rightarrow \boxed{\text{C}}$

Answer: (C) [Go Back to Q4](#)

Q5.

Solution

Concept – Poisson probability: $P(X = k) = \frac{e^{-\lambda} \lambda^k}{k!}$, with mean rate λ .

Step 1 – list the parameters: $\lambda = 2$ per metre, $k = 0$.



Step 2 – substitute into the formula: $P(X = 0) = \frac{e^{-2} 2^0}{0!}$

Step 3 – evaluate the power and factorial: $2^0 = 1$ and $0! = 1$

$$P(X = 0) = e^{-2}$$

Step 4 – substitute $e^{-2} = 0.1353$: $P(X = 0) = 0.1353$

Final Answer: $P(X = 0) = 0.1353 \Rightarrow \boxed{\text{D}}$

Answer: (D) [Go Back to Q5](#)

Q6.

Solution

Concept – expectation of a function of X : $E[X^2] = \sum_i x_i^2 P(x_i)$.

Step 1 – square each value: $(-1)^2 = 1$, $0^2 = 0$, $1^2 = 1$.

Step 2 – weight each square by its probability: $1(0.3) + 0(0.4) + 1(0.3)$

Step 3 – evaluate each term: $= 0.3 + 0 + 0.3$

Step 4 – add: $E[X^2] = 0.6$

Final Answer: $E[X^2] = 0.6 \Rightarrow \boxed{\text{C}}$

Answer: (C) [Go Back to Q6](#)

Q7.

Solution

Concept – margin of error of a confidence interval: For a known- σ mean, the margin of error is $E = z \frac{\sigma}{\sqrt{n}}$.

Step 1 – compute $\sqrt{n}$: $\sqrt{64} = 8$

Step 2 – compute the standard error: $\frac{\sigma}{\sqrt{n}} = \frac{16}{8} = 2$

Step 3 – multiply by the critical value: $E = 1.96 \times 2$

Step 4 – evaluate: $E = 3.92$

Step 5 – interpret: The 95% interval is 40 ± 3.92 , i.e. (36.08, 43.92); its half-width is 3.92.

Final Answer: margin of error = 3.92 $\Rightarrow \boxed{\text{D}}$

Answer: (D) [Go Back to Q7](#)



Q8.

Solution

Concept – Pearson correlation coefficient: $r = \frac{S_{xy}}{\sqrt{S_{xx} S_{yy}}}$.

Step 1 – list the sums of squares: $S_{xx} = 25$, $S_{yy} = 100$, $S_{xy} = 30$.

Step 2 – compute the product under the root: $S_{xx} S_{yy} = 25 \times 100 = 2500$

Step 3 – take the square root: $\sqrt{2500} = 50$

Step 4 – form the ratio: $r = \frac{30}{50}$

Step 5 – simplify: $r = 0.60$

Why the other options are wrong: 1.20 is the slope $S_{xy}/S_{xx} = 30/25$, not the correlation, and $|r|$ can never exceed 1.

Final Answer: $r = 0.60 \Rightarrow$ D

Answer: (D) [Go Back to Q8](#)

Q9.

Solution

Concept – binomial probability: $P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}$.

Step 1 – list the parameters: $n = 5$, $k = 2$, $p = 0.5$, $1 - p = 0.5$.

Step 2 – evaluate the binomial coefficient: $\binom{5}{2} = \frac{5 \times 4}{2 \times 1} = 10$

Step 3 – combine the powers: $p^2(1 - p)^3 = (0.5)^2(0.5)^3 = (0.5)^5$

Step 4 – evaluate the power: $(0.5)^5 = \frac{1}{32} = 0.03125$

Step 5 – multiply: $P(X = 2) = 10 \times 0.03125 = 0.3125$

Final Answer: $P(X = 2) = 0.3125 \Rightarrow$ C

Answer: (C) [Go Back to Q9](#)

Q10.

Solution

Concept – exponential cumulative probability: For an exponential variable with rate λ , $P(X \leq t) = 1 - e^{-\lambda t}$.

Step 1 – find the rate from the mean: Mean = $\frac{1}{\lambda} = 4$ seconds, so $\lambda = \frac{1}{4} = 0.25 \text{ s}^{-1}$.



Step 2 – form the exponent: $\lambda t = 0.25 \times 4 = 1$

Step 3 – write the cumulative probability: $P(X \leq 4) = 1 - e^{-1}$

Step 4 – substitute $e^{-1} = 0.3679$: $P(X \leq 4) = 1 - 0.3679 = 0.6321$

Final Answer: $P(X < 4) \approx 0.632 \Rightarrow \boxed{\text{B}}$

Answer: (B) [Go Back to Q10](#)

Q11.

Solution

Concept – rank is the number of independent rows: Reduce the rows and count how many are linearly independent.

Step 1 – compare rows 2 and 1: Row 2 = (2, 4, 2) = $2 \times (1, 2, 1) = 2R_1$, so row 2 depends on row 1.

Step 2 – compare rows 3 and 1: Row 3 = (3, 6, 3) = $3 \times (1, 2, 1) = 3R_1$, so row 3 also depends on row 1.

Step 3 – eliminate the dependent rows: $R_2 \leftarrow R_2 - 2R_1 = (0, 0, 0)$ and $R_3 \leftarrow R_3 - 3R_1 = (0, 0, 0)$.

Step 4 – count the non-zero independent rows: Only R_1 remains non-zero, giving rank = 1.

Final Answer: $\text{rank}(A) = 1 \Rightarrow \boxed{\text{A}}$

Answer: (A) [Go Back to Q11](#)

Q12.

Solution

Concept – cofactor expansion along the first row: $\det B = b_{11}M_{11} - b_{12}M_{12} + b_{13}M_{13}$.

Step 1 – first-row entries: $b_{11} = 1, b_{12} = 2, b_{13} = 0$.

Step 2 – minor for b_{11} : $\begin{vmatrix} 3 & 1 \\ 0 & 1 \end{vmatrix} = (3)(1) - (1)(0) = 3$

Step 3 – minor for b_{12} : $\begin{vmatrix} 0 & 1 \\ 2 & 1 \end{vmatrix} = (0)(1) - (1)(2) = -2$

Step 4 – the third term vanishes: Since $b_{13} = 0$, its contribution is 0.

Step 5 – combine with the sign pattern: $\det B = 1(3) - 2(-2) + 0 = 3 + 4$



$$= 7$$

Final Answer: $\det B = 7 \Rightarrow$ B

Answer: (B) [Go Back to Q12](#)

Q13.

Solution

Concept – eigenvalues from trace and determinant: For a 2×2 matrix, $\lambda^2 - \text{tr}(C)\lambda + \det(C) = 0$.

Step 1 – trace and determinant of C : $\text{tr}(C) = 2 + 3 = 5$ and $\det(C) = (2)(3) - (2)(1) = 6 - 2 = 4$.

Step 2 – write the characteristic equation: $\lambda^2 - 5\lambda + 4 = 0$

Step 3 – factor: $(\lambda - 1)(\lambda - 4) = 0$

Step 4 – read the roots: $\lambda = 1$ or $\lambda = 4$.

Step 5 – pick the smaller: The smaller eigenvalue is 1.

Final Answer: $\lambda_{\min} = 1 \Rightarrow$ B

Answer: (B) [Go Back to Q13](#)

Q14.

Solution

Concept – Frobenius norm via the trace: $\|A\|_F = \sqrt{\text{tr}(A^T A)}$, and the trace equals the sum of the eigenvalues.

Step 1 – sum the eigenvalues of $A^T A$: $\text{tr}(A^T A) = 16 + 9 = 25$

Step 2 – take the square root: $\|A\|_F = \sqrt{25}$

Step 3 – evaluate: $\|A\|_F = 5$

Why the other options are wrong: 25 is the trace itself, before taking the root; the singular values are $\sqrt{16} = 4$ and $\sqrt{9} = 3$, whose squares sum to 25.

Final Answer: $\|A\|_F = 5 \Rightarrow$ A

Answer: (A) [Go Back to Q14](#)



Q15.

Solution

Concept – the ℓ_2 norm is the square root of the sum of squares: $\|v\|_2 = \sqrt{\sum_i v_i^2}$.

Step 1 – square each component: $6^2 = 36$ and $8^2 = 64$.

Step 2 – add the squares: $36 + 64 = 100$

Step 3 – take the square root: $\|v\|_2 = \sqrt{100} = 10$

Step 4 – contrast with the other norms: $\|v\|_1 = 6 + 8 = 14$ and $\|v\|_\infty = \max(6, 8) = 8$; the ℓ_2 value is 10.

Final Answer: $\|v\|_2 = 10 \Rightarrow \boxed{\text{B}}$

Answer: (B) [Go Back to Q15](#)

Q16.

Solution

Concept – minimize by setting the derivative to zero: For a convex parabola, the stationary point is the global minimum.

Step 1 – differentiate: $f'(x) = 2x - 6$

Step 2 – set the derivative to zero: $2x - 6 = 0$

$x = 3$

Step 3 – confirm it is a minimum: $f''(x) = 2 > 0$, so the point is a minimum.

Step 4 – evaluate f at $x = 3$: $f(3) = 3^2 - 6(3) + 13$

$= 9 - 18 + 13$

$= 4$

Final Answer: minimum value $= 4 \Rightarrow \boxed{\text{C}}$

Answer: (C) [Go Back to Q16](#)

Q17.

Solution

Concept – scalar projection: The component of $\vec{a}$ along $\vec{b}$ is $\text{comp}_{\vec{b}} \vec{a} = \frac{\vec{a} \cdot \vec{b}}{\|\vec{b}\|}$.

Step 1 – compute the dot product: $\vec{a} \cdot \vec{b} = (3)(4) + (4)(0) = 12 + 0 = 12$

Step 2 – compute the magnitude of $\vec{b}$: $\|\vec{b}\| = \sqrt{4^2 + 0^2} = \sqrt{16} = 4$



Step 3 – form the scalar projection: $\text{comp}_{\vec{b}} \vec{a} = \frac{12}{4}$

Step 4 – divide: $\text{comp}_{\vec{b}} \vec{a} = 3$

Step 5 – interpret: Since $\vec{b}$ points along the x -axis, the projection is just the x -component of $\vec{a}$, which is 3.

Final Answer: scalar projection = 3 $\Rightarrow$ **B**

Answer: (B) [Go Back to Q17](#)

Q18.

Solution

Concept – one gradient-descent step: $x_1 = x_0 - \eta f'(x_0)$.

Step 1 – differentiate the objective: $f(x) = (x - 3)^2 \Rightarrow f'(x) = 2(x - 3)$

Step 2 – evaluate the gradient at $x_0 = 0$: $f'(0) = 2(0 - 3) = -6$

Step 3 – apply the update with $\eta = 0.2$: $x_1 = 0 - 0.2 \times (-6)$

Step 4 – finish the arithmetic: $x_1 = 0 + 1.2 = 1.2$

Step 5 – sanity check the direction: The minimum is at $x = 3$; moving from 0 towards 1.2 heads in the correct direction.

Final Answer: $x_1 = 1.2 \Rightarrow$ **A**

Answer: (A) [Go Back to Q18](#)

Q19.

Solution

Concept – convexity via the Hessian: A twice-differentiable function is convex iff its Hessian is positive semi-definite everywhere.

Step 1 – compute the second partials: $f_{xx} = 4, f_{yy} = 6, f_{xy} = f_{yx} = 1$.

Step 2 – write the Hessian: $H = \begin{bmatrix} 4 & 1 \\ 1 & 6 \end{bmatrix}$

Step 3 – test with the leading principal minors: First minor: $4 > 0$. Second minor: $\det H = (4)(6) - (1)(1) = 24 - 1 = 23 > 0$.

Step 4 – conclude: Both leading minors are positive, so H is positive definite and f is convex everywhere.

Final Answer: f is convex $\Rightarrow$ **C**

Answer: (C) [Go Back to Q19](#)



Q20.

Solution

Concept – extended (step) slicing in Python: `a[start:stop:step]` takes every step-th element; `a[: :2]` means start at 0, go to the end, step 2.

Step 1 – index the list: $a = \left[\underbrace{1}_0, \underbrace{2}_1, \underbrace{3}_2, \underbrace{4}_3, \underbrace{5}_4 \right]$

Step 2 – apply the step of 2 from index 0: Take indices 0, 2, 4.

Step 3 – read off the elements: Index 0 $\rightarrow$ 1, index 2 $\rightarrow$ 3, index 4 $\rightarrow$ 5, giving [1, 3, 5].

Final Answer: [1, 3, 5] $\Rightarrow$ B

Answer: (B) [Go Back to Q20](#)

Q21.

Solution

Concept – reversing a sequence with a negative step: `x[::-1]` walks the sequence from the last element to the first, producing the reverse.

Step 1 – list the characters with their indices: h(0), e(1), l(2), l(3), o(4)

Step 2 – apply the step of -1: Start at the last index 4 and move backwards: 4, 3, 2, 1, 0.

Step 3 – read off the characters in that order: o, l, l, e, h

Step 4 – join into a string: The result is olleh.

Final Answer: olleh $\Rightarrow$ A

Answer: (A) [Go Back to Q21](#)

Q22.

Solution

Concept – a queue is first-in, first-out (FIFO): `dequeue()` always removes the element that was enqueued earliest.

Step 1 – enqueue(1), enqueue(2): Queue (front $\rightarrow$ rear): [1, 2].

Step 2 – dequeue(): Removes the front 1; queue becomes [2].

Step 3 – enqueue(3), enqueue(4): Queue becomes [2, 3, 4].

Step 4 – dequeue(): Removes the front 2; queue becomes [3, 4].

Step 5 – dequeue() (the last one): Removes the front 3; this is the value returned.



Final Answer: last dequeue() returns 3 $\Rightarrow$

Answer: (C) [Go Back to Q22](#)

Q23.

Solution

Concept – a leaf has no children: Build the BST and count nodes with no descendants.

Step 1 – place the root: 15 is inserted first and becomes the root.

Step 2 – place the next two keys: $10 < 15$ goes left; $20 > 15$ goes right.

Step 3 – place the remaining keys: $8 < 15$ then $8 < 10$ (left of 10); $12 < 15$ then $12 > 10$ (right of 10).

Step 4 – identify the leaves: The nodes with no children are 8, 12, 20.

Step 5 – count them: There are 3 leaf nodes.

Final Answer: 3 leaves $\Rightarrow$

Answer: (B) [Go Back to Q23](#)

Q24.

Solution

Concept – linear probing scans the next free slot: On a collision at index i , try $i + 1, i + 2, \dots$ (modulo table size).

Step 1 – insert 15: $h(15) = 15 \bmod 7 = 1$; slot 1 is free, so 15 goes to index 1.

Step 2 – insert 22: $h(22) = 22 \bmod 7 = 1$; index 1 is occupied, so probe index 2; it is free, so 22 goes to index 2.

Step 3 – insert 8: $h(8) = 8 \bmod 7 = 1$; index 1 occupied $\rightarrow$ probe index 2 (occupied) $\rightarrow$ probe index 3 (free).

Step 4 – place 8: 8 is stored at index 3.

Final Answer: index 3 $\Rightarrow$

Answer: (C) [Go Back to Q24](#)



Q25.

Solution

Concept – binary search halves the range each step: The worst-case number of comparisons is $\lfloor \log_2 n \rfloor + 1$.

Step 1 – bracket $n = 63$ between powers of 2: $2^5 = 32$ and $2^6 = 64$, so $32 < 63 < 64$.

Step 2 – take the base-2 logarithm: $\log_2 63 \approx 5.98$

Step 3 – take the floor: $\lfloor 5.98 \rfloor = 5$

Step 4 – add one: $5 + 1 = 6$

Final Answer: 6 comparisons $\Rightarrow$ **C**

Answer: (C) [Go Back to Q25](#)

Q26.

Solution

Concept – Merge Sort recurrence: Merge Sort always splits the array in half and merges in linear time, giving $T(n) = 2T(n/2) + \Theta(n)$.

Step 1 – identify the recurrence parameters: $a = 2$, $b = 2$, $f(n) = \Theta(n)$.

Step 2 – compute the critical exponent: $n^{\log_b a} = n^{\log_2 2} = n$

Step 3 – compare $f(n)$ with $n^{\log_b a}$: $f(n) = \Theta(n) = \Theta(n^{\log_b a})$, which is Master-theorem Case 2.

Step 4 – apply Case 2: $T(n) = \Theta(n \log n)$

Why the other options are wrong: Unlike Quicksort, Merge Sort's split is always balanced, so it never degrades to $O(n^2)$; its worst case is still $O(n \log n)$.

Final Answer: $O(n \log n) \Rightarrow$ **D**

Answer: (D) [Go Back to Q26](#)

Q27.

Solution

Concept – DFS explores as deep as possible before backtracking: Use a stack (or recursion); always follow the smallest unvisited neighbour first.

Step 1 – adjacency lists (sorted): $A: B, C; \quad B: A, D; \quad C: A, D, F; \quad D: B, C, E; \quad E: D; \quad F: C.$

Step 2 – start at A, go to B: Visit A, then its smallest neighbour B.



Step 3 – from B , go to D : B 's unvisited neighbour is D . Visit D .

Step 4 – from D , go to C : D 's smallest unvisited neighbour is C (since $C < E$). Visit C .

Step 5 – from C , go to F : C 's unvisited neighbour is F . Visit F ; it has no new neighbour, so backtrack.

Step 6 – backtrack to D , go to E : D 's remaining unvisited neighbour is E . Visit E .

Step 7 – assemble the order: A, B, D, C, F, E .

Final Answer: $A, B, D, C, F, E \Rightarrow \boxed{D}$

Answer: (D) [Go Back to Q27](#)

Q28.

Solution

Concept – Master theorem: For $T(n) = aT(n/b) + f(n)$, compare $f(n)$ with $n^{\log_b a}$.

Step 1 – identify the parameters: $a = 4, b = 2, f(n) = n$.

Step 2 – compute the critical exponent: $n^{\log_b a} = n^{\log_2 4} = n^2$

Step 3 – compare $f(n)$ with $n^{\log_b a}$: $f(n) = n = O(n^{2-\varepsilon})$ for $\varepsilon = 1$, so f is polynomially smaller: this is Case 1.

Step 4 – apply Case 1: $T(n) = \Theta(n^{\log_b a}) = \Theta(n^2)$

Final Answer: $T(n) = \Theta(n^2) \Rightarrow \boxed{A}$

Answer: (A) [Go Back to Q28](#)

Q29.

Solution

Concept – size bounds of a set union: For union-compatible relations, $\max(|R|, |S|) \leq |R \cup S| \leq |R| + |S|$; the maximum occurs when the two sets are disjoint.

Step 1 – read the tuple counts: $|R| = 8, |S| = 5$.

Step 2 – disjoint case gives the maximum: If no tuple is shared, every tuple survives the union.

Step 3 – add: $|R \cup S|_{\max} = 8 + 5 = 13$

Why the other options are wrong: $40 = 8 \times 5$ is the Cartesian-product size, not a union; 8 is the *minimum* (when $S \subseteq R$).



Final Answer: 13 tuples $\Rightarrow$

Answer: (A) [Go Back to Q29](#)

Q30.

Solution

Concept – COUNT(DISTINCT c) counts distinct non-NULL values: It removes duplicates and ignores NULLs.

Step 1 – list the column values: 10, 10, 20, NULL, 30, 20.

Step 2 – drop the NULLs: The NULL is ignored, leaving 10, 10, 20, 30, 20.

Step 3 – keep only distinct values: The distinct values are {10, 20, 30}.

Step 4 – count them: There are 3 distinct non-null values.

Final Answer: COUNT(DISTINCT c) returns 3 $\Rightarrow$

Answer: (C) [Go Back to Q30](#)

Q31.

Solution

Concept – 3NF versus BCNF: 3NF allows a dependency $X \rightarrow Y$ if Y is a prime attribute; BCNF requires the left side X to be a superkey.

Step 1 – identify prime attributes: The candidate keys are AB and AC , so the prime attributes are A , B and C (all three).

Step 2 – check 2NF (partial dependency): No non-prime attribute exists, so there can be no partial dependency; R is in 2NF.

Step 3 – check 3NF for $C \rightarrow B$: B is a prime attribute, so the dependency $C \rightarrow B$ is allowed under 3NF. R is in 3NF.

Step 4 – check BCNF for $C \rightarrow B$: C alone is not a superkey (its closure $C^+ = \{B, C\}$ does not include A), so $C \rightarrow B$ violates BCNF.

Step 5 – conclude the highest normal form: R is in 3NF but not BCNF, so the highest normal form is 3NF.

Final Answer: highest form = 3NF $\Rightarrow$

Answer: (B) [Go Back to Q31](#)



Q32.

Solution

Concept – natural join matches tuples on the common attribute: Combine a tuple of R with a tuple of S whenever their B values are equal.

Step 1 – match R 's tuple $(1, p)$: S has (p, x) and (p, y) with $B = p$, giving $(1, p, x)$ and $(1, p, y)$ – 2 tuples.

Step 2 – match R 's tuple $(2, p)$: S again has (p, x) and (p, y) with $B = p$, giving $(2, p, x)$ and $(2, p, y)$ – 2 tuples.

Step 3 – match R 's tuple $(3, q)$: S has (q, z) with $B = q$, giving $(3, q, z)$ – 1 tuple.

Step 4 – total the matches: $2 + 2 + 1 = 5$

Final Answer: 5 tuples $\Rightarrow$

Answer: (D) [Go Back to Q32](#)

Q33.

Solution

Concept – OLAP cube operations: Slice fixes one dimension to a single value; dice selects a sub-cube on several dimensions; roll-up and drill-down move between aggregation levels.

Step 1 – describe the operation: Fixing Year = 2023 keeps only that one value of the Year dimension.

Step 2 – effect on dimensionality: Holding one dimension at a single value drops it from the result, cutting the cube's dimensionality by one.

Step 3 – match the name: Choosing a single value on one dimension is a **slice**.

Why the other options are wrong: Dice picks a range on *several* dimensions (dimensionality unchanged); roll-up and drill-down change the level of aggregation, not the number of dimensions.

Final Answer: slice $\Rightarrow$

Answer: (D) [Go Back to Q33](#)



Q34.

Solution

Concept – order and keys in a B^+ tree: An internal node of order m holds at most m child pointers, and the number of keys is always one fewer than the number of pointers.

Step 1 – maximum pointers: Order $m = 7 \Rightarrow$ up to 7 child pointers.

Step 2 – keys are one fewer than pointers: Maximum keys $= m - 1 = 7 - 1$

Step 3 – evaluate: $= 6$

Final Answer: 6 keys $\Rightarrow$ A

Answer: (A) [Go Back to Q34](#)

Q35.

Solution

Concept – attribute closure: A^+ is the set of all attributes functionally determined by A , found by repeatedly applying the FDs.

Step 1 – start with the attribute itself: $A^+ = \{A\}$

Step 2 – apply $A \rightarrow B$: A determines B , so $A^+ = \{A, B\}$.

Step 3 – apply $B \rightarrow C$: B is now in the closure, so it brings in C : $A^+ = \{A, B, C\}$.

Step 4 – check for more: No FD has D or E on its right-hand side reachable from these attributes, so nothing further is added.

Final Answer: $A^+ = \{A, B, C\} \Rightarrow$ A

Answer: (A) [Go Back to Q35](#)

Q36.

Solution

Concept – logistic regression prediction: Compute the linear score $z = w^T x + b$, then apply the sigmoid.

Step 1 – compute the linear score: $z = (2)(1) + (-1)(2) + 0$

$$= 2 - 2 + 0$$

$$= 0$$

Step 2 – write the sigmoid: $\sigma(0) = \frac{1}{1 + e^{-0}}$

Step 3 – evaluate e^{-0} : $e^{-0} = 1$



Step 4 – finish: $\sigma(0) = \frac{1}{1+1} = \frac{1}{2} = 0.50$

Final Answer: $\sigma(0) = 0.50 \Rightarrow \boxed{D}$

Answer: (D) [Go Back to Q36](#)

Q37.

Solution

Concept – underfitting in bias–variance terms: A model too simple to capture the pattern has high bias; if it is also insensitive to the training sample, it has low variance.

Step 1 – interpret the low training accuracy: Failing to fit even the training data means the model is too rigid, i.e. *high bias*.

Step 2 – interpret the similar test accuracy: Performance barely changes across data sets, meaning it does not chase sample noise, i.e. *low variance*.

Step 3 – combine: This is the classic signature of underfitting: high bias, low variance.

Final Answer: high bias, low variance $\Rightarrow \boxed{A}$

Answer: (A) [Go Back to Q37](#)

Q38.

Solution

Concept – 1-NN takes the label of the single closest point: Compute the Euclidean distance from Q to each training point and pick the nearest.

Step 1 – distance $Q(4, 4)$ to $A(1, 1)$: $\sqrt{(4-1)^2 + (4-1)^2} = \sqrt{9+9} = \sqrt{18} = 4.24$

Step 2 – distance to $B(5, 5)$: $\sqrt{(4-5)^2 + (4-5)^2} = \sqrt{1+1} = \sqrt{2} = 1.41$

Step 3 – distance to $C(2, 2)$: $\sqrt{(4-2)^2 + (4-2)^2} = \sqrt{4+4} = \sqrt{8} = 2.83$

Step 4 – distance to $D(6, 6)$: $\sqrt{(4-6)^2 + (4-6)^2} = \sqrt{4+4} = \sqrt{8} = 2.83$

Step 5 – pick the nearest point: The smallest distance is 1.41, to B , which is Class 1.

Step 6 – assign the label: With $k = 1$, Q takes B 's label, Class 1.

Final Answer: Q is Class 1 $\Rightarrow \boxed{A}$

Answer: (A) [Go Back to Q38](#)



Q39.

Solution

Concept – binary entropy: $H = -p \log_2 p - (1 - p) \log_2 (1 - p)$, with p the fraction of positives.

Step 1 – compute the class proportions: $p = \frac{5}{10} = 0.5$ positive, $1 - p = \frac{5}{10} = 0.5$ negative.

Step 2 – take the base-2 logs: $\log_2 0.5 = -1$ for both terms.

Step 3 – form each term: $-0.5 \times (-1) = 0.5$

$$-0.5 \times (-1) = 0.5$$

Step 4 – add the terms: $H = 0.5 + 0.5 = 1.0$ bit

Step 5 – interpret: A perfectly balanced binary set has the maximum possible entropy of 1 bit.

Final Answer: $H = 1.0$ bit $\Rightarrow$ **B**

Answer: (B) [Go Back to Q39](#)

Q40.

Solution

Concept – information gain: $IG = H(\text{parent}) - \sum_{\text{children}} \frac{n_{\text{child}}}{n_{\text{parent}}} H(\text{child})$.

Step 1 – parent entropy: The parent has 4+ and 4– (balanced), so $H(\text{parent}) = 1.0$ bit.

Step 2 – entropy of the $A = 0$ child: It has 4+ and 0–; a pure node has entropy 0.

Step 3 – entropy of the $A = 1$ child: It has 0+ and 4–; a pure node has entropy 0.

Step 4 – weighted child entropy: $\frac{4}{8}(0) + \frac{4}{8}(0) = 0$

Step 5 – subtract from the parent entropy: $IG = 1.0 - 0 = 1.0$

Step 6 – interpret: The attribute separates the classes perfectly, so the gain equals the full parent entropy.

Final Answer: $IG = 1.0$ bit $\Rightarrow$ **A**

Answer: (A) [Go Back to Q40](#)



Q41.

Solution

Concept – SVM margin width: For the hyperplane $w^T x + b = 0$, the full margin (distance between the two boundaries) is $\frac{2}{\|w\|}$.

Step 1 – write the margin formula: $\text{margin} = \frac{2}{\|w\|}$

Step 2 – substitute $\|w\| = 2$: $\text{margin} = \frac{2}{2}$

Step 3 – divide: $\text{margin} = 1$

Final Answer: margin width = 1 $\Rightarrow$ **D**

Answer: (D) [Go Back to Q41](#)

Q42.

Solution

Concept – k-means centroid update: The new centroid is the coordinate-wise mean of the points assigned to the cluster.

Step 1 – average the x -coordinates: $\bar{x} = \frac{2 + 4 + 6}{3} = \frac{12}{3} = 4$

Step 2 – average the y -coordinates: $\bar{y} = \frac{4 + 4 + 10}{3} = \frac{18}{3} = 6$

Step 3 – assemble the centroid: $(\bar{x}, \bar{y}) = (4, 6)$

Why the other options are wrong: (12, 18) is the coordinate sum, not the mean; the centroid divides each sum by the 3 points.

Final Answer: centroid = (4, 6) $\Rightarrow$ **B**

Answer: (B) [Go Back to Q42](#)

Q43.

Solution

Concept – variance explained in PCA: Each principal component explains variance proportional to its eigenvalue; the fraction is the component's eigenvalue over the total.

Step 1 – total variance (sum of all eigenvalues): $6 + 2 + 1 + 1 = 10$

Step 2 – variance in the first component: The largest eigenvalue is 6.

Step 3 – form the fraction: $\frac{6}{10} = 0.60$



Step 4 – express as a percentage: $0.60 = 60\%$

Final Answer: 60% of the variance $\Rightarrow$

Answer: (D) [Go Back to Q43](#)

Q44.

Solution

Concept – neuron forward pass with ReLU: Compute the weighted sum plus bias, then apply $\text{ReLU}(z) = \max(0, z)$.

Step 1 – weighted sum of inputs: $w_1x_1 + w_2x_2 = (1)(2) + (2)(1)$

$$= 2 + 2$$

$$= 4$$

Step 2 – add the bias: $z = 4 + (-1) = 3$

Step 3 – apply ReLU: $\text{ReLU}(3) = \max(0, 3) = 3$

Final Answer: output = 3 $\Rightarrow$

Answer: (B) [Go Back to Q44](#)

Q45.

Solution

Concept – A* expands the node of least $f = g + h$: Compute f for each frontier node and pick the minimum.

Step 1 – f for node P : $f(P) = g + h = 4 + 3 = 7$

Step 2 – f for node Q : $f(Q) = 2 + 6 = 8$

Step 3 – f for node R : $f(R) = 5 + 1 = 6$

Step 4 – choose the smallest f : $\min(7, 8, 6) = 6$, which belongs to node R .

Final Answer: A* expands $R \Rightarrow$

Answer: (D) [Go Back to Q45](#)



Q46.

Solution

Concept – minimax back-up: MIN nodes take the minimum of their children; the MAX root takes the maximum of the MIN results.

Step 1 – evaluate the left MIN node: $\min(6, 4) = 4$

Step 2 – evaluate the right MIN node: $\min(8, 5) = 5$

Step 3 – evaluate the MAX root: $\max(4, 5) = 5$

Final Answer: minimax value = 5 $\Rightarrow$ **C**

Answer: (C) [Go Back to Q46](#)



Answer Key

Q	Ans	Q	Ans	Q	Ans	Q	Ans	Q	Ans
1	C	2	A	3	B	4	C	5	D
6	C	7	D	8	D	9	C	10	B
11	A	12	B	13	B	14	A	15	B
16	C	17	B	18	A	19	C	20	B
21	A	22	C	23	B	24	C	25	C
26	D	27	D	28	A	29	A	30	C
31	B	32	D	33	D	34	A	35	A
36	D	37	A	38	A	39	B	40	A
41	D	42	B	43	D	44	B	45	D
46	C								

